



Technological Coupling and Ethical Reconstruction: A Systematic Inquiry into the Legal Paradigm Shift in the AI Era

Bingying Zhou

Associate Professor, School of Economics and Management, Guangzhou Institute of Science and Technology, Guangdong, China

chow6117@alu.ruc.edu.cn

Received: 16 Feb 2026; Received in revised form: 18 Mar 2026; Accepted: 22 Mar 2026; Available online: 28 Mar 2026

Abstract—The leapfrog evolution of Artificial Intelligence (AI) is profoundly reshaping the operational logic of legal practice and the normative dimensions of legal ethics. Drawing upon cutting-edge scholarship from 2019 to 2025, this paper systematically examines the paradigm shift at the intersection of law and AI through three critical lenses: technological empowerment, ethical conflicts, and regulatory pathways. Technologically, Generative AI (GenAI) has transitioned from a mere auxiliary tool to a pivotal "actor" within legal practice. While the integration of Retrieval-Augmented Generation (RAG) and associated technologies has mitigated model hallucinations and enhanced algorithmic explainability, the inherent risks posed by technological boundaries remain substantial. Ethically, GenAI exerts systemic pressure on attorneys' duties of competence and confidentiality. By analyzing judicial precedents such as the UK High Court's *Ayinde* case, this study underscores the urgency of reconstructing professional responsibility and advocates for a recalibrated ethical framework through the lens of the "inevitability of bias." Regulatorily, the "legislative-driven" model of the EU, the "market-led" approach of the US, and the "consensus-building" strategy of international organizations exhibit distinct governance characteristics, whereas the concept of "co-regulation" provides a theoretical foundation for coupling technical standards with legal frameworks. Research indicates that the legal order in the AI era is at a critical juncture, transitioning from "heterogeneous isomorphism" toward "organic integration." Future scholarship should prioritize empirical validation, deepen cross-jurisdictional comparisons, and facilitate the collaborative design of technical solutions and ethical principles, thereby fostering a holistic paradigm update in legal education and institutional design.

Keywords—Generative AI, Legal Ethics, Co-regulation, Legal Paradigm, Retrieval-Augmented Generation (RAG), Professional Responsibility.

I. INTRODUCTION: PROBLEM STATEMENT AND RESEARCH FRAMEWORK

Over the past five years, interdisciplinary research on AI and law has undergone a paradigm shift from

"marginal instrumental exploration" to "core normative agenda." This transformation is driven not only by the application potential unleashed by technological breakthroughs such as Large Language

Models (LLMs) but also by the systemic impact of AI practices on existing legal theories and ethical frameworks. From the *Mata v. Avianca* case in 2023, where a lawyer mistakenly submitted fictitious case law generated by ChatGPT, to the *Ayinde* case in 2025, which revealed the serious consequences of GenAI fabricating legal authorities, these incidents of "technological anomie" in judicial practice have evolved from isolated professional mistakes into landmark events in legal ethics discussions. These events prompt a fundamental inquiry in legal scholarship: "When AI systems become deeply embedded in the daily operations of legal practice, do traditional professional conduct norms and regulatory frameworks face a 'regulatory deficit'? In the technologically mediated legal field, how should the boundaries of human actors' responsibility be defined, and how should the legal order respond to the quasi-autonomous characteristics exhibited by technological systems?" In response, legal theory has engaged in profound normative reflection. Mei Xiaying [1] points out that AI legislation should be built upon the new types of legal objects and methods it creates, proposing "deep learning algorithms" and "human-machine relationships" as two theoretical pillars of AI law. This insight suggests that the relationship between AI and law is not a simple mechanical overlay of "technological tools + legal regulation" but a comprehensive paradigm shift involving ontology, epistemology, and methodology. Ontologically, AI systems as novel "actors" in legal processes challenge the human-centered subjectivity structure of classical jurisprudence. Epistemologically, legal practitioners urgently need to construct a cognitive framework capable of penetrating the logic of technological operations. Methodologically, traditional doctrinal legal research must achieve deep integration of knowledge resources with disciplines such as computer science and cognitive science.

Based on cutting-edge research from Chinese and international scholars between 2019 and 2025, this paper reveals key issues and theoretical frontiers in the intersection of AI and law through in-depth

exploration and systematic review. By comparatively analyzing and theoretically integrating the technical aspects of GenAI, the ethical challenges it poses, and the regulatory pathways, we can grasp more profoundly its evolutionary logic and capability boundaries within the legal field, recognize more clearly its systemic impact on lawyers' ethical duties and the ethical warnings from judicial practice, and conduct more robust comparative analysis and theoretical integration of regulatory paths. Through a systematic exploration of the transformation of the legal order in the AI era, this paper not only provides an analytical framework with both theoretical depth and practical relevance for understanding the formation of the new legal ecology in this complex era but also identifies potential directions and ideas for future research.

II. TECHNOLOGICAL SOLUTIONS: FROM AUXILIARY TOOL TO LEGAL "ACTOR"

2.1 The Technological Evolution of Legal AI

Early legal AI primarily relied on symbolism and expert systems, consequently encountering the "engineering bottleneck" of knowledge acquisition and the "fragmentation" of reasoning. The technological logic of this phase was to formalize legal rules and reasoning processes into computable knowledge representations. The core difficulty lay in the fact that the openness, context-dependence, and value-laden nature of legal knowledge could not be fully encoded. For instance, rule-based expert systems required legal experts and knowledge engineers to work together to translate legal norms into IF-THEN logical rules. However, the interpretation of legal concepts often involves considerations of purpose, value balancing, and policy judgments—elements difficult to fully express within a closed rule system.

In recent years, LLMs, represented by the Transformer architecture, have overcome the limitations of traditional methods through capabilities like contextual reasoning, few-shot learning, and generative argumentation. LLMs employ statistical learning methods to "capture"

statistical regularities and semantic associations from vast amounts of text data, enabling them to generate natural language text tailored to specific contextual requirements. This approach of automatically extracting legal writing styles, argumentation patterns, and reasoning paths from data, driven by the "hard" force of data, effectively avoids the difficulty of manually encoding complex legal knowledge into formal rules and significantly reduces the rigidity and limitations inherent in "hard" legal rules. The advantage of this technological path lies in its ability to automatically learn legal writing styles, argumentation patterns, and reasoning paths through data-driven methods, eliminating the need for manual encoding of legal knowledge into formal rules.

The first comprehensive survey of legal LLMs, jointly released by multiple Chinese and international institutions [2], proposed an innovative "dual-perspective taxonomy." This survey integrates classical legal reasoning theories like Toulmin's argumentation framework, mapping structural elements of legal argument—claim, data, warrant, qualifier, rebuttal—onto the functional design of AI systems. Simultaneously, it maps the authentic workflows of professional roles such as lawyers, judges, and litigants, enabling technological development to respond to the embodied needs of legal practice. By systematically reviewing nearly 60 tools and datasets covering key technological advancements in legal text processing, knowledge integration, and reasoning formalization, the study provides a theoretical blueprint for the transition of legal AI from "laboratory tool" to "judicial infrastructure."

2.2 Application Landscape of GenAI in Legal Practice

Research by Terzidou [3] indicates that GenAI systems have become deeply embedded in the daily operations of legal practice: from electronic communication and document drafting to contract template generation, multilingual translation, and long-text summarization; LLMs are reshaping lawyers' workflows. In legal research, GenAI has been integrated into mainstream platforms like

LexisNexis and Westlaw, offering legal practitioners rapid access to sources and streamlined drafting. These platforms utilize natural language processing, allowing lawyers to pose legal questions conversationally, with the system returning relevant cases, statutes, and secondary sources, and generating preliminary legal analysis suggestions. The application of GenAI is particularly significant in contract analysis and drafting. Lawyers can input contract texts into the system, requesting identification of potential risk clauses, suggestions for revisions, or generation of specific types of contract templates. Based on learning from vast amounts of contract text, AI systems can identify industry-standard clauses, common negotiation points, and ambiguous phrasing that might lead to disputes. This application not only improves the efficiency of contract review but also helps standardize contract quality and reduce human error. In litigation support, GenAI is used for evidence summarization, case fact organization, and drafting legal memoranda. Lawyers can input large volumes of evidence, asking the system to generate timelines, identify key facts, or summarize party arguments. For complex litigation requiring massive document processing, AI support can significantly reduce lawyers' workload, allowing them to focus on strategic judgment and argument innovation.

Abbasi [4] conceptualizes this trend as "Responsible Legal Augmentation," emphasizing that AI systems should be oriented towards augmenting rather than replacing human judgment. This reflects a philosophical reflection on legal augmentation and awareness of potential ethical dilemmas. He advocates orienting AI's value towards augmenting, not replacing, human judgment. The distinction between AI as "automation" versus "assistance" directly relates to the intensity of responsibility attribution and ethical constraints: if AI is seen as a "automation" tool replacing humans, responsibility might be "outsourced" to AI, reducing ethical constraints; if AI is seen as an "assistance" system augmenting human capabilities, human ultimate control and oversight obligations must be retained.

Abbasi argues that responsible legal augmentation should satisfy three core requirements: transparency (lawyers understand AI's capabilities and limitations), auditability (AI outputs can be effectively reviewed), and accountability (humans always bear ultimate responsibility for final decisions).

2.3 Technological Capability Boundaries and Core Dilemmas

Despite the promising applications of GenAI in law, core dilemmas stemming from its technological capability boundaries persist. The "hallucination" problem in LLMs' legal assertions constitutes a primary challenge. LLMs may generate entirely fictitious cases, statutes, or arguments with a high degree of confidence, and the reliance of legal reasoning on authoritative sources makes this issue particularly dangerous. The cause of model hallucinations lies in the LLM's training objective: to generate coherent text fitting the context, not to ensure factual accuracy. When the model encounters legal questions not covered in its training data, it tends to "fill in the gaps" through plausible guesswork rather than admitting ignorance. As explicitly stated in the UK High Court's *Ayinde* judgment, freely available LLM-based GenAI tools (like ChatGPT) cannot be relied upon for reliable legal research; the critical safeguard is "verifying any output through authoritative sources." This warning is universally applicable: no matter how advanced AI systems become, lawyers bear ultimate responsibility for the accuracy of their submissions. The court's stance can be summarized as "responsibility reinforcement under technological neutrality": not presupposing the virtue or vice of specific technologies, but consistently reinforcing the ultimate control and oversight obligations of human actors.

Second, the adaptability gap in low-resource jurisdictions is also noteworthy. LLMs' training data predominantly comes from legal texts from the English-speaking world. Model performance often significantly declines for non-English jurisdictions and legal traditions outside common law systems.

This issue concerns not only technological fairness but also the normative value of legal pluralism. Profound differences exist in legal concepts, reasoning methods, and argumentation styles across jurisdictions. Applying a model trained on data from one jurisdiction directly to another may lead to systematic misunderstandings and erroneous reasoning.

Furthermore, the lack of explainability in black-box reasoning constitutes another core dilemma. The internal workings of LLMs involve hundreds of millions or even hundreds of billions of parameters, making the reasoning path leading to a specific output difficult to fully trace. For legal decisions, explainability is not merely a technical requirement but the foundation of legitimacy – litigants have the right to know the basis of decisions, and judges have the duty to state their reasoning. When AI systems participate in legal reasoning, ensuring that their contribution can be understood, reviewed, and challenged becomes a pressing methodological issue.

To address these challenges, Retrieval-Augmented Generation (RAG) technology significantly enhances the accuracy and timeliness of generated content by invoking external real-time knowledge bases rather than relying solely on static training data. The RAG architecture combines the generative capabilities of LLMs with the retrieval capabilities of external knowledge bases: when a user asks a question, the system first retrieves relevant documents from the knowledge base and then inputs these documents as context to the LLM, generating an answer based on the retrieved results. This technological path transforms the production and updating of legal knowledge from traditional "closed model training" to "open knowledge retrieval," offering new possibilities for the explainability and accountability of legal AI. To truly establish effective linkage between AI answers and human control, AI responses must have traceable authority, enabling lawyers to independently verify the "basis" relied upon by the AI, thereby establishing an effective

connection between technical assistance and human control.

III. LEGAL ETHICS: RECONSTRUCTION OF PROFESSIONAL RESPONSIBILITY

3.1 The Systemic Impact of GenAI on Lawyers'

Ethical Duties

The core duties of competence and confidentiality for lawyers face systemic impacts due to the involvement of GenAI. A research report by Terzidou [3] deeply analyzes the specific manifestations and coping strategies for this impact. The research indicates that lawyers' attitudes towards GenAI show significant generational and firm-size differences: younger lawyers and those in large firms are more inclined to actively experiment with AI tools, while senior lawyers and those in smaller firms are more cautious. This disparity highlights the urgency and the need to respect diverse perspectives in reconstructing legal professional ethics.

Under the duty of competence, lawyers can effectively use GenAI systems by acquiring necessary digital literacy to enhance the efficiency, consistency, and quality of case handling. This has two aspects: positively, lawyers should understand available AI tools and their capabilities, utilizing them in appropriate scenarios to fulfill their commitment to providing competent representation to clients; negatively, lawyers must recognize the limitations of these systems and avoid relying on their outputs where their validity, accuracy, or relevance is questionable. Relevant opinions from the American Bar Association's ethics committee explicitly state that lawyers' duty to review AI-generated content is no different from their review duty in traditional legal research—ultimate responsibility always rests with the signing lawyer.

The reconstruction of competence also involves redefining the standard of "competence." In the AI era, does ignorance of available technological tools constitute professional negligence? Does failure to effectively use AI assistance mean failing to provide the best possible service to clients? These questions lack uniform answers, but it is clear that the meaning

of competence is expanding from "mastery of legal knowledge" to "mastery of all practice resources, including technological tools." Lawyers need critical skills to evaluate AI outputs and balance technical assistance with independent judgment.

The duty of confidentiality faces a more direct challenge. Lawyers inputting client information directly into AI prompts or allowing AI systems to access files containing confidential data may lead to client information disclosure to third parties (including system providers), thereby violating the duty of confidentiality. Research shows that the scope of confidentiality protection should cover electronic communications between lawyer and client, including metadata in electronic documents. This means lawyers must carefully evaluate the data storage and processing policies of commercial AI tools to prevent client information from being used without authorization for model training or other commercial purposes.

Terzidou's research [3] also reveals that many lawyers lack understanding of AI systems' data processing practices, mistakenly assuming interactions with AI are as confidential as discussions with colleagues. This cognitive bias poses a significant risk: commercial AI systems often use user-input data for model optimization, and client information may become part of the training data without the lawyer's knowledge, potentially even being indirectly exposed in other users' queries. Therefore, when selecting AI tools, lawyers must review the service provider's data protection policies to ensure client confidentiality is adequately protected.

3.2 Ethical Warnings from Judicial Practice: From the Mata Case to the Ayinde Case

The 2023 *Mata v. Avianca* case in the US, where lawyers submitted fictitious case law generated by ChatGPT, has become a landmark event in legal ethics discussions. In that case, lawyers relied on ChatGPT for legal research but failed to verify the generated content, resulting in a brief containing multiple non-existent judicial precedents. The court's sanctions emphasized lawyers' duty of candor to the

court—regardless of technological development, lawyers are ultimately responsible for their submissions. Notably, the lawyers involved did not intend to deceive the court but lacked awareness of AI systems' capability boundaries, mistakenly believing ChatGPT could conduct reliable legal research. The core lesson is that technological reliance cannot substitute for professional judgment; lawyers must maintain skeptical scrutiny of AI outputs.

The UK High Court's *Ayinde* case further deepened the discussion. The judgment established an important ethical principle: lawyers cannot defend false citations by claiming the "legal principle itself is correct." Citations are not merely decorative or optional; they are the bedrock of legal reasoning and judicial trust. Even if the legal principle supported by a cited case is correct, the presence of a fictitious case still undermines the integrity of the judicial process because it prevents opposing parties and the court from verifying the basis of arguments, destroying the common foundation of legal dialogue. The UK High Court's stance in this case reflects a prudent and pragmatic approach: neither banning the use of GenAI as AI skeptics might desire nor unconditionally endorsing its reliability, but establishing a middle path by emphasizing lawyers' duty to the court—ensuring all submitted materials are accurate and not misleading. This stance can be summarized as "responsibility reinforcement under technological neutrality": courts do not presuppose the virtue or vice of specific technologies but consistently reinforce the ultimate control and oversight obligations of human actors. The guidance particularly emphasizes that lawyers must fulfill a verification duty when using AI tools; any statement relying on AI-generated content must be validated through authoritative sources.

3.3 Bias and Value: A Counterintuitive Insight

Lande [5] offers a counterintuitive theoretical perspective: "bias" in AI should not be simplistically viewed as a technical flaw to be eradicated but rather as an inherent and potentially positive feature of system operation. This view stems from the premise that technology is never neutral. The notion of

"value-neutral" technology is an illusion; any system is a product of specific design choices and value judgments. From dataset selection and algorithm feature setting to the presentation of final results, every technical detail is permeated with developers' subjective presuppositions and beliefs. Based on this, Lande argues that bias reflects developers' values, and instead of trying to hide or erase these positions, they should be disclosed. This shifts the discussion from "how to eliminate bias" to "how to responsibly disclose bias." The research suggests that enabling users to understand the assumptions, priorities, and design frameworks behind a tool should be a core ethical guideline. For example, with a legal research tool, if it preferentially cites cases from specific courts or arguments with specific logic without prior explanation, users might mistakenly believe they have received a comprehensive, neutral answer. Once these preferences are clearly disclosed, users can understand that the result is from a specific perspective.

Lande further categorizes bias into two types: harmful bias and constructive bias. Harmful bias leads to discrimination against specific groups or systematically distorts factual truth. Constructive bias reflects legitimate value choices, such as designing systems to prioritize human rights, emphasize procedural justice, or pursue operational efficiency. The focus is not on pursuing the impossible goal of "zero bias" but on making these value orientations transparent, serving as the basis for users to make informed choices.

Zhou, S. [6] deepens this discussion from the perspective of virtue jurisprudence. The research points out that the term "responsible AI" carries a misleading anthropomorphic tendency that easily blurs true responsibility attribution. It argues for a strict distinction between "legal responsibility" and "moral capacity." Legally, responsibility can recognize AI as an "actor" with attributes of legal responsibility, making the legal effects it generates traceable. Morally, it insists on a "human-centered" baseline: AI cannot become a subject of moral responsibility. This shift has important practical

implications: we should not focus on whether "AI is responsible" but rather on "whether the humans using AI possess the virtues to operate technology responsibly." This redirects the ethical focus from the cold algorithm back to humans' own decision-making capabilities.

3.4 The Complex Landscape of Responsibility Attribution

With GenAI's involvement in legal practice, the issue of responsibility attribution has acquired unprecedented complexity. When an AI system generates erroneous output causing client loss, who should bear responsibility? The lawyer using the AI, the company developing the AI, the law firm deploying the AI, or the AI system itself? Answering this question involves multiple areas like contract law, tort law, and professional responsibility law, with no unified theoretical framework currently available.

From a contract law perspective, legal service contracts between lawyers and clients may implicitly or explicitly stipulate the use of AI. If a lawyer uses AI to process sensitive information without client consent, it may constitute a breach of contract; if a lawyer promises personal service but excessively relies on AI, it may violate the essence of the service contract. From a tort law perspective, a lawyer's reliance on AI outputs may constitute negligence, requiring elements like breach of duty of care, damages, and causation. From a product liability perspective, AI systems might be considered products, and developers might be liable for system defects, but the learning capabilities and autonomy of AI challenge the traditional product liability framework.

"consensus-building" model. These three models reflect fundamental differences in legal philosophies towards technology governance and embody distinct political traditions and institutional resources.

The EU model, centered on the Artificial Intelligence Act (AI Act), creatively introduces a "harmonized standards" mechanism, making compliance with technical standards an important basis for proving that AI systems meet mandatory legal requirements, thus achieving close coupling between technical standards and the legal framework. Yousefi [8] provides an in-depth interpretation of the EU AI Act from the perspective of algorithmic fairness in their doctoral thesis, emphasizing that the Act adopts a risk-based approach, classifying AI systems into four risk levels: unacceptable risk, high risk, limited risk, and minimal or no risk, with different regulatory requirements applicable to each. High-risk AI systems (e.g., used in recruitment, credit assessment, predictive policing) must meet strict compliance requirements, including risk management systems, data governance, technical documentation, transparency obligations, and human oversight. The EU model's advantage lies in providing a deterministic legal framework, offering clear behavioral expectations for businesses and users; its risk is that legislation may lag behind technological development, and rigid norms may stifle innovation. To mitigate this tension, the EU AI Act introduces "regulatory sandboxes," allowing innovators to test new technologies under regulatory supervision, seeking balance between compliance and innovation.

The US model emphasizes market leadership and industry self-regulation. The National Institute of Standards and Technology (NIST) issued the Artificial Intelligence Risk Management Framework (AI RMF), serving primarily as a guiding tool to help designers, developers, and users of AI systems manage risks associated with AI applications. The institutional logic of this model is that technological innovation outpaces legislative processes; overly rigid legal regulation may dampen industrial vitality. Through industry self-regulation and soft law

IV. REGULATORY PATHWAYS: COMPARATIVE LAW AND INSTITUTIONAL DESIGN

4.1 Three Regulatory Models in a Global Perspective

Globally, AI regulation presents a landscape of coexisting diverse models. A systematic study by Zhang Tao [7] compares three typical paths: the EU's "legislative-driven" model, the US's "market-led" model, and the international organizations'

governance, risk control can be achieved while maintaining flexibility. The US model's advantage is high flexibility and adaptability, allowing quick responses to technological changes; its risk is insufficient regulatory force, potentially leading to market failures, consumer rights infringement, and accumulation of systemic risks. In recent years, discussions on AI regulation in the US have become increasingly active, with some states already enacting specific AI legislation and federal-level legislative initiatives advancing.

International standardization organizations (such as ISO/IEC JTC 1/SC 42) focus on building global technical consensus, having published dozens of international standards in areas like concepts and terminology, risk management, governance impact, and trustworthiness. The advantage of this model lies in transcending sovereign state boundaries, providing a common technical language and behavioral expectations for global industrial chains; its limitations include the democratic deficit in standard-setting processes and the weakness of enforcement mechanisms. International standard development relies primarily on consensus among technical experts, with insufficient participation from affected groups (e.g., consumers, vulnerable populations); adoption of standards depends on voluntary choices by states, lacking mandatory enforcement mechanisms.

4.2 Theoretical Construction of Co-regulation

Facing the coexistence and competition of diverse regulatory models, Zhang Tao [7] proposes that compared to purely administrative regulation models or pure self-regulation models, a "co-regulation model" is more aligned with the inherent laws of AI standardization and can more effectively address the practical dilemmas of standardization.

The theoretical foundations of the co-regulation model include embedded ethics, governance ecology, knowledge co-creation, and layered governance. Embedded ethics integrates social ethical values into the technology design process, transforming ethical considerations from external constraints into internal constituent elements. This approach requires ethicists

and technologists to collaborate in the early stages of technology development, translating value demands like fairness, transparency, and explainability into specific technical requirements. Haden [9] deepens this approach from a human rights-centered perspective, advocating drawing on the institutional experience of the General Data Protection Regulation (GDPR) to construct an AI regulatory framework centered on human rights protection. She advocates extending Data Protection Impact Assessments to AI Human Rights Impact Assessments, anticipating potential human rights risks during the system design phase and taking mitigating measures.

Governance ecology emphasizes the network effects of diverse stakeholders. AI regulation involves multiple actors—developers, deployers, users, regulators, affected individuals and groups—and the knowledge and capacity of any single actor are limited. The co-regulation model integrates distributed knowledge and coordinates diverse interests by constructing a governance ecology. This approach requires establishing regular stakeholder dialogue mechanisms so that concerns of different actors can be expressed and considered in the regulatory process.

Knowledge co-creation posits that regulatory knowledge is not monopolized by legislators or regulators but is dynamically generated through the interaction of multiple actors. This perspective breaks the cognitive hierarchy of traditional "command-and-control" regulation, providing a theoretical basis for the coexistence of diverse normative forms like soft law, standards, and guidelines. In the context of rapid technological iteration, knowledge co-creation means the regulatory framework must remain open and adaptive, capable of absorbing cutting-edge experience from practice.

Layered governance distinguishes regulatory needs at different levels. The technical infrastructure layer (e.g., algorithm architecture), the application layer (e.g., specific scenarios), and the social impact layer (e.g., human rights protection) require different types of regulatory tools and institutional

arrangements; layered governance helps achieve precise matching of regulatory tools. For example, regulation at the technical infrastructure layer can focus on safety and reliability standards; application-level regulation can focus on risk assessment and mitigation for specific scenarios; social impact-level regulation needs to consider broader value balance and rights protection.

4.3 Chinese Practice and Institutional Exploration

In recent years, China has successively issued policy documents such as the New Generation Artificial Intelligence Ethics Norms and the Opinions on Strengthening Science and Technology Ethics Governance, initially constructing a normative framework for AI governance. In the field of standardization, China actively participates in international standard-setting while promoting the construction of its domestic standard system, issuing important documents like the White Paper on Artificial Intelligence Standardization. Zhang Tao [7] notes that China's regulatory path exhibits characteristics of "policy guidance + standard support." At the policy level, the state defines development goals and value orientations through strategic planning; at the standard level, technical standards provide operational tools for policy implementation. The advantage of this path lies in the combination of policy flexibility and standard professionalism, capable of absorbing the knowledge contributions of technical experts while maintaining strategic direction. However, China's AI regulation also faces unique challenges. How to ensure citizen rights while promoting technological innovation, how to enhance voice in international standard-setting, and how to coordinate regulatory functions across different government departments—these questions require ongoing institutional exploration and theoretical reflection.

4.4 Regulatory Exploration in the Judicial Sphere

Besides legislation and administrative regulation, proactive regulation of AI within the judicial sphere is also noteworthy. The guidance issued by the UK High Court in the *Ayinde* case reflects the judiciary's cautious approach towards AI use: neither

prohibiting its use as AI skeptics might expect nor unconditionally endorsing it, but establishing a middle path by emphasizing lawyers' duty to the court to ensure all submitted materials are accurate and not misleading.

This judicial regulatory path is characterized by its focus on concrete behavioral norms in specific scenarios rather than abstract presuppositions about technology's virtue or vice; it does not create new substantive rules but activates existing ethical duties in new technological contexts; it does not rely on rigid prohibitions or permissions but guides behavior through the clarification of responsibility attribution. This model of "regulation through activating duties" offers important methodological insights for legal governance during periods of rapid technological iteration. Judicial regulation is also evident in courts' determinations of the admissibility of AI-generated evidence. When evidence generated by AI systems (e.g., deepfake detection reports, algorithmic risk assessments) is submitted to court, judges must assess its reliability, relevance, and fairness. This process is simultaneously an application of evidence rules and a judicial review of technical norms, providing a practical testing ground for technical standards for AI systems.

V. RESEARCH OUTLOOK: TOWARD A NEW JURISPRUDENCE FOR THE AI ERA

5.1 Limitations of Current Research and Future Agenda

While significant progress has been made in current interdisciplinary research on AI and law, several directions warrant further development: First, the expansion of empirical research. Existing studies are mostly normative analyses, lacking empirical evaluation of the effects of AI systems in actual judicial operations. For example, what is the actual accuracy rate of AI-assisted decision-making? Do different groups show varying levels of trust in AI decisions? Does AI use change judges' sentencing patterns? Does AI assistance truly improve the efficiency and quality of legal services? These questions require support from large-scale empirical

research, necessitating interdisciplinary collaboration combining law, computer science, sociology, psychology, and other fields. Second, the deepening of cross-jurisdictional comparisons. There are significant differences in the acceptance and regulatory paths for AI across different legal systems and jurisdictions with varying levels of development. Comparative research can provide references for global governance. Particularly noteworthy is how differences between common law and civil law systems—regarding precedent status, reasoning methods, procedural structures—affect the integration paths and regulatory needs for AI technology. Additionally, gaps between developing and developed countries in technological capacity, data resources, and regulatory capability require attention in comparative research. Third, the collaborative design of technical solutions and ethical principles. How to embed ethical principles like fairness, transparency, and explainability into technical architecture, achieving "ethics by design," is a core issue for technology-law interdisciplinary research. This requires not only enhanced ethical awareness among technologists but also basic technological literacy among legal scholars to intervene during the technology design phase rather than offering post-hoc criticism. Collaborative design necessitates establishing interdisciplinary cooperation mechanisms to translate legal insights into constraints and design goals for technical implementation. Fourth, the adaptive transformation of legal education. The innovative practice of the AI course on Legal Professional Ethics at China University of Political Science and Law demonstrates that what competency structure do legal professionals need in the AI era, and how to integrate technological literacy and ethical awareness into legal education—these are not only educational issues but also fundamental questions of professional ethics. The course design team's proposed "knowledge-problem-competence" three-part framework and their attempt to construct highly simulated case scenarios using AI agents provide valuable experience for the digital transformation of

legal education. Future legal education must incorporate emerging competencies like technological literacy, data analysis, and human-AI collaboration alongside traditional legal knowledge, cultivating a new generation of legal professionals capable of practicing effectively in an AI environment.

5.2 From Heterogeneous Isomorphism to Organic Integration: The Transformation of the Legal Paradigm

Mei Xiaying [1] points out that AI possesses distinct characteristics such as systematicity, publicity, openness, and evolutionary nature, and its integration with the traditional legal system forms a relationship of "heterogeneous isomorphism." "Heterogeneous isomorphism" implies that AI and the traditional legal system are nested within each other in structural form but have fundamental differences in operational logic and value core. This judgment has important methodological significance for understanding the transformation of jurisprudence in the AI era. Moving from "heterogeneous isomorphism" to "organic integration" requires paradigm updates at several levels.

At the ontological level, the boundaries of legal subjects need re-examination. When AI systems can autonomously generate texts with legal effects, participate in legal reasoning processes, and influence judicial outcomes, the traditional human-centered theory of legal subjects faces challenges. However, this does not mean granting AI legal personality but rather acknowledging its status as an "actor" while maintaining the human-centered baseline for responsibility attribution. This position requires conceptual distinction between the "source of legal effects" and the "attribution of legal responsibility"—the former may involve AI systems, but the latter always belongs to humans.

At the epistemological level, cognitive frameworks for understanding technological systems need development. Legal professionals require basic technological literacy to understand the capability boundaries and operational logic of LLMs, avoiding mystifying or demonizing technology. Legal

scholarship needs to shift from "legal responses to technology" to "legal cognition co-evolving with technology." This means legal methodology must incorporate emerging cognitive models like computational thinking and systems thinking, enabling legal judgment to be grounded in an understanding of technology's actual operation.

At the methodological level, the integration of diverse research paradigms is necessary. Research at the intersection of AI and law is inherently interdisciplinary, requiring the integration of knowledge resources from law, computer science, ethics, cognitive science, and other fields. The methodological framework of a single discipline cannot capture the complex picture of technology-law-society interactions. Interdisciplinary research is not simple accumulation but requires establishing shared problem awareness, conceptual frameworks, and research methods, enabling knowledge from different disciplines to generate new insights through dialogue.

5.3 Conclusion: Safeguarding the Rule of Law Amidst the Technological Wave

The transformative evolution of AI technology is reshaping every aspect of legal practice, from legal research and document drafting to case prediction and adjudication assistance, from law practice to court management. This process brings both opportunities for efficiency gains and knowledge democratization and profound challenges such as the reconstruction of professional ethics, blurred responsibility attribution, and human rights protection risks.

Facing this transformation, we should neither succumb to the blind optimism of technological utopia nor fall into the pessimistic resistance of technological dystopia. The UK High Court's prudent stance in the Ayinde case perhaps offers a valuable attitude: neither presupposing the virtue or vice of specific technologies nor avoiding the challenges they bring; neither abandoning vigilance towards technological boundaries nor abandoning exploration of technological potential; neither relaxing the emphasis on human actors' responsibility nor

denying the factual status of AI systems as new "actors."

Ultimately, the rule of law order in the AI era remains a human-centered order. Technology can enhance human capabilities, expand human cognition, and optimize human decisions, but it cannot replace human judgment, human responsibility, or human dignity. Safeguarding the rule of law amidst the technological wave means embracing technological possibilities while steadfastly upholding the core values of the rule of law: human autonomy, attributability of responsibility, and availability of remedy for rights. This is both the generational mission of legal professionals and the enduring topic of legal scholarship.

Research at the intersection of AI and law is in an active phase of knowledge innovation. Over the next decade, as technology continues to evolve and regulatory experience accumulates, this field will undoubtedly yield richer theoretical achievements and practical wisdom. This review is merely a starting point; more scholars are expected to join this important academic dialogue, collectively exploring the shape of the legal order in the AI era.

REFERENCES

- [1] Mei, Xiaying. (2025). Fundamental Issues in AI Legislation: Algorithmic Discipline and Human-Machine Relationships in Context. *Jurist* (6), 16-30. (in Chinese)
- [2] Shao, P., Xu, L., Wang, J., Zhou, W., & Wu, X. (2025). When large language models meet law: Dual-lens taxonomy, technical advances, and ethical governance. *arXiv preprint arXiv:2507.07748*.
- [3] Terzidou, K. (2025). Generative ai systems in legal practice offering quality legal services while upholding legal ethics. *International Journal of Law in Context*, 21(3).
- [4] Abbasi, Z. (2025). Responsible legal augmentation: integrating generative ai into legal practice. *Law, Technology & Humans*, 7(3).
- [5] Lande J. (2025). RPS Coach is "Biased" - And Proud of It [R]. University of Missouri School of Law Legal

- Studies Research Paper No. 2025-17,
2025. <https://dx.doi.org/10.2139/ssrn.5213572>
- [6] Zhou, S. (2025). Deconstructing 'Responsible AI': An Examination of Legal and Ethical Accountability Through Virtue Jurisprudence. *International Journal for the Semiotics of Law-Revue internationale de Sémiotique juridique*, 1-22.
- [7] Zhang, Tao. (2025). Regulating Artificial Intelligence through Technical Standards: A Jurisprudence Based on Co-regulation. *Comparative Law Studies* (4), 169-187. (in Chinese)
- [8] Yousefi, Y. (2025). The quest to fairness in algorithmic decisions: ethical, legal and technical solutions. [D]. ORBilu-University of Luxembourg. <https://orbilu.uni.lu/handle/10993/66101>, 2025.
- [9] Haden, C. (2025). Assembling the Algorithm for Human Rights Centred AI Regulation (Doctoral dissertation, University of Hertfordshire).
- [10] Custers, B., & Fosch-Villaronga, E. (2022). Law and artificial intelligence. *Information technology and law series*, 35.
- [11] DiMatteo, L. A., Poncibò, C., & Cannarsa, M. (Eds.). (2022). *The Cambridge handbook of artificial intelligence: global perspectives on law and ethics*. Cambridge University Press.